

Placation and provocation

Rationality and Society
2014, Vol. 26(1) 73–104

© The Author(s) 2014

Reprints and permissions:
sagepub.co.uk/journalsPermissions.nav

DOI: 10.1177/1043463113513092

rss.sagepub.com



Joshua Tasoff

Claremont Graduate University, USA

Abstract

It has been observed that industries self-regulate to placate a regulator from taking action, and revolutionary vanguards sometimes provoke an apathetic populace into revolt. This paper presents a very simple model that captures this strategic maneuvering, and applies it to several other examples. Two players have preferences over the realization of a policy; the first player has a marginal cost to affect the policy and the second player has a fixed cost. The fixed cost provides strategic incentives for the first mover to placate or provoke the second player. In equilibrium, the second mover may benefit from having preferences that diverge more from the first mover, and may benefit by having higher fixed costs. As the number of first movers increases, placation and provocation both become more likely, and the second player's incentives to occlude or reveal its fixed cost become stronger as well.

Keywords

Fixed costs, provision point, revolution, self-regulation

Introduction

There are many instances in which industries self-regulate in a way that seems to go against their profit interests. The film industry in the US self-censors its own films through the Motion Picture Association of America (MPAA). There is an analogous organization in the video game industry, and there are many other self-regulating organizations across many

Corresponding author:

Joshua Tasoff, Department of Economics, 160 E. Tenth Street, Claremont, CA 91711, USA.

Email: joshua.tasoff@cgu.edu

industries, such as the American Medical Association in health and the Financial Industry Regulatory Authority in finance. But whereas some of these self-regulating bodies presumably exist to solve market failures stemming from asymmetric information, it is less clear why a self-regulating body would exist in the case of the MPAA, where there do not exist any clear agency problems.

The historical record offers a fairly clear story of the motivation of the industry. In the 1920s films became more morally ambiguous relative to the social values of the time. Citizens protested and federal censorship loomed over the film industry. Knowing that there were substantial fixed costs in establishing a censorship agency, the film industry placated Congress by instead self-regulating (Ellis and Wexman, 2002). The restriction on their own production was a means to deter even harsher restrictions from Congress.

This event has features in common with a completely unrelated social interaction—political revolutions. It is often said that revolutionaries strive to worsen conditions for the populace in order to spur them to action. In fact, several revolutionaries have expressed exactly this sentiment in political manifestos, including Argentine Marxist revolutionary Che Guevara and Brazilian Marxist revolutionary Carlos Marighela (Marighela, 1972). The populace has major fixed costs to organizing and revolting because there is a tremendous collective action problem. Thus the revolutionary vanguard may worsen the political condition for the populace in order to spark a revolt. Here, the vanguard goes against the seeming interests of the population that it purports to fight for in order to provoke action from the populace that ultimately leads to an outcome in the direction of the vanguard's interest.

The common logic in both these situations is that the second mover, the government or the populace, has a fixed cost to taking action. The action moves the “policy,” a variable that both players care about. In the film example the policy is the “moral” permissiveness of film standards, and in the revolutionary example the policy is a measure of the political environment. Both the film industry and the revolutionaries move the policy in the opposite direction of their interest for strategic reasons, due to the second mover's fixed costs. In the former case, the industry moves the state variable so that the fixed costs deter government action. In the latter case, the revolutionaries move the policy so that it is optimal for the populace to incur the fixed costs in order to take action.

This strategic “back-pedaling” is a consequence of the second mover's fixed costs and it will occur depending on the relationship between the first

and second mover's preferences, as well as their marginal costs to take action. This can be described formally with a simple game. Two agents have single-peaked preferences over a single policy dimension. The agents play an extensive-form game where the first player has a marginal cost associated with moving the policy and the second player has a fixed cost. The implication of the fixed cost in this simple context is that it creates an "inaction zone"—whenever the policy is sufficiently close to its bliss point, the second player prefers the costless action. Under certain conditions, this structure provides strategic incentives for the first player to move the policy such that it barely placates the second player from taking action, or barely provokes the second player to take action.

The novel predictions are as follows:

1. The second mover may be better off when its preferences diverge from the first mover. For example, the populace may be better off when its preferences diverge from the revolutionaries'. When preferences are sufficiently distant, the second mover will be spared the first mover's provocation because the second mover's desired outcome is too distant from the first mover's. This contrasts with the intuition that greater discrepancy in players' preferences results in lower expected utility for both players.
2. The second mover may prefer to possess greater fixed costs in order to dissuade the first mover from provocation. By doing so, the second mover strategically ties its hands as a form of commitment to deter provocation from the first mover.
3. When the model is extended with multiple first movers that move simultaneously, both the placation and provocation equilibrium regimes become more stable. The logic is that the placation and provocation regimes provide a "provision point" such that all first movers are pivotal in inducing the desired response from the second mover. The alternative regimes, on the other hand, suffer from a free-riding problem that is exacerbated as the number of first movers grows large. Thus self-regulation becomes more likely as the number of firms in an industry increases.
4. As the number of first movers increases, the second mover's incentive to occlude or reveal its fixed costs increases. In the case of placation, revelation becomes more effective at encouraging placation with many first movers. For example, the regulating body has more incentive to reveal its fixed costs when the industry has many firms. In the case of provocation, occlusion becomes more effective as a deterrence with many first movers.

This is not the first paper on self-regulation. Several scholars have observed the strategic incentives of self-regulation (DeMarzo et al., 2005; Maxwell et al., 2000; Stango, 2003; Stefanadis, 2003; Volden and Wiseman, 2008), and those involved in inciting revolution (Acemoglu and Robinson, 2001; Grossman, 1991; Roemer, 1985; Tullock, 1971). The primary purpose of this paper is to present a very simple conceptual lens through which to view numerous social interactions. The model is not designed for mathematical generality or applicative specificity, but rather as a very simple structure to help understand and express the commonalities across a wide variety of strategic situations. In addition, understanding why and when the model does not apply can help to uncover fundamental aspects of a strategic interaction.

The next section formally presents the model and establishes some basic results. Numerous applications of the model are explored in the third section. The fourth section explores extensions of the model and the fifth section concludes.

Model

The game is an extensive-form game with two players, $i \in \{1, 2\}$, who have a single-peaked strictly concave felicity function, $f_i(\cdot)$, over a single policy dimension q . At the beginning of the game the default policy is $\bar{q}$. Player 1 acts first and can choose to move the policy an amount, $m_1 \in \mathbb{R}$. Then Player 2 may move the policy $m_2 \in \mathbb{R}$. The realized policy, $q = \bar{q} + m_1 + m_2$, determines the utility of both players.

The bliss point is defined as the realized policy q , that maximizes the felicity function. Without loss of generality, Player 2 has bliss point 0 and Player 1 has bliss point b , where it is assumed $b \geq 0$. The parameter b is the bias of Player 1 relative to Player 2. When b is small, the bliss points of both players are near, and when b is large the bliss points are distant.

Both players have a weakly convex cost function that is positive and monotonically increasing in the magnitude of Player i 's action, $|m_i|$, and where initial costs and initial marginal costs are zero, $C_i(0)=0$ and $C'_i(0)=0$. Furthermore, we assume the cost function is increasing in the absolute value of the action so that $C'_i(x) > 0$ when $x > 0$, and $C'_i(x) < 0$ when $x < 0$. The main difference between the two utility functions is that Player 2 has a fixed cost of e to moving the policy a nonzero distance, $m_2 \neq 0$. Both felicity and cost functions are assumed to be differentiable as many times as necessary. The utility functions are given by:

$$U_i(m_1, m_2) = f_i(\bar{q} + m_1 + m_2) - C_i(m_i) \quad (1)$$

$$U_2(m_1, m_2) = f_2(\bar{q} + m_1 + m_2) - C_2(m_2) - e * 1\{m_2 \neq 0\} \quad (2)$$

Notice that if Player 2 chooses to pay the fixed costs associated with taking an action, then Player 2 will move the realized policy q such that it maximizes $U_2(m_1, m_2)$. However, if this gain in utility is not worth the fixed cost of taking an action, then no action will be taken. This fixed cost produces an inaction zone around Player 2's bliss point. Let q_L and q_U denote the lower and upper bounds of the inaction zone respectively. Define the inaction zone as the interval $[q_L, q_U]$ where if $m_1 + \bar{q} \in [q_L, q_U]$, then Player 2's optimal action is $m_2 = 0$. Thus, Player 2's bliss point is in the inaction zone, $0 \in [q_L, q_U]$. The subgame perfect equilibrium will depend on the location of the two bliss points, the span and location of the inaction zone, and the location of the default policy.

Equilibrium characterization and regimes

Since this is an extensive game with perfect information, the equilibrium concept employed is subgame perfect equilibrium. Generically, the parameters uniquely determine the pure-strategy subgame perfect equilibrium. We categorize the equilibria into four equilibrium regimes, each of which has an equilibrium path that is qualitatively similar. These regimes each exist in a region of the $\bar{q}$, b , and e parameter space. The four regimes are named placation, pull, provocation, and improvement and are graphed in $\bar{q}$ - b space in Figure 1, and graphed in $\bar{q}$ - e space in Figure 2. An alternative perspective is presented in Figure 3 which depicts the equilibrium regimes in b - e space. These four qualitatively different equilibrium regimes can be illustrated with the two motivating examples.

Self-regulation: Placation and pull

The motivating example for the placation logic is the story of the American film industry's self-regulation in the early twentieth century. In 1915, the Supreme Court case *Mutual Film Corporation v. Industrial Commission of Ohio* ruled that motion pictures were not covered by the First Amendment and that the government had the power to censor film (Ellis and Wexman, 2002). The threat of federal regulation loomed over the industry. The film industry's solution to this problem was to form its own censorship entity. By self-regulating, it could restrict itself minimally while still placating the government. A film history textbook explains (Mast and Kavin, 1992):

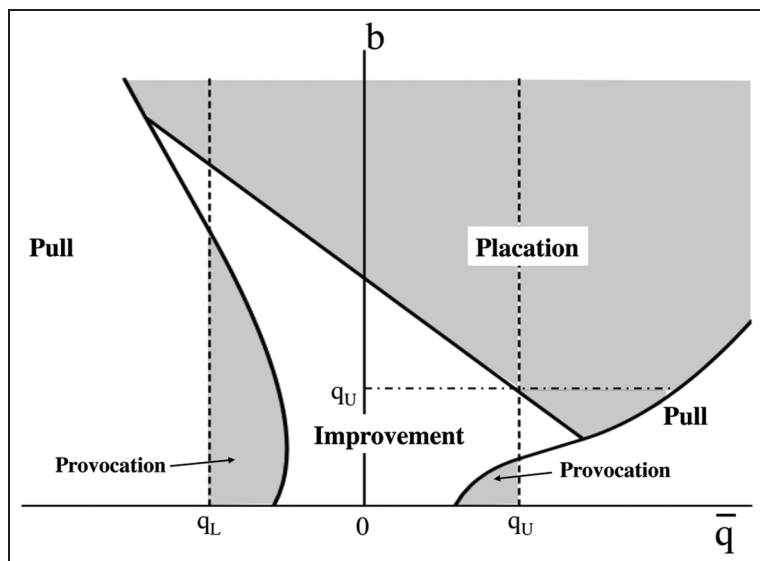


Figure 1. Equilibrium regions in $\bar{q}$ – b space.

Note: In this graph $C_2(\cdot) = 0$.

Such notoriety brought the film business to the attention of the United States Congress and to the edge of federal censorship – the last thing any producer wanted. The industry decided once again to clean its own house, to serve as its own censorship body ... The loose informal advising of the Hays Office in the 1920s was another in a series of successful Hollywood attempts to keep films out of the hands of government censors.

The model operationalizes this example: the film industry moves first and the government follows. The government has fixed costs, $e > 0$, to establishing a censorship bureau. Both care about the level of “iniquity” q in film. A high level of “iniquity” means more profits for the industry (since sex and violence sell) but upsets the sensibilities of those in government. So, in other words, the bliss point of the industry is higher than the bliss point of government. The film industry discourages the government from taking any action by voluntarily lowering the level of “iniquity” $m_1 = q_U - \bar{q}$ so that the policy is on the edge of government’s inaction zone q_U . Government takes no action $m_2 = 0$. Ultimately, this leaves the total level of violence and sex higher than if government had taken action. This is labeled as “placation” in Figure 1. In a placation equilibrium Player 1 moves the policy away from its bliss point, resulting

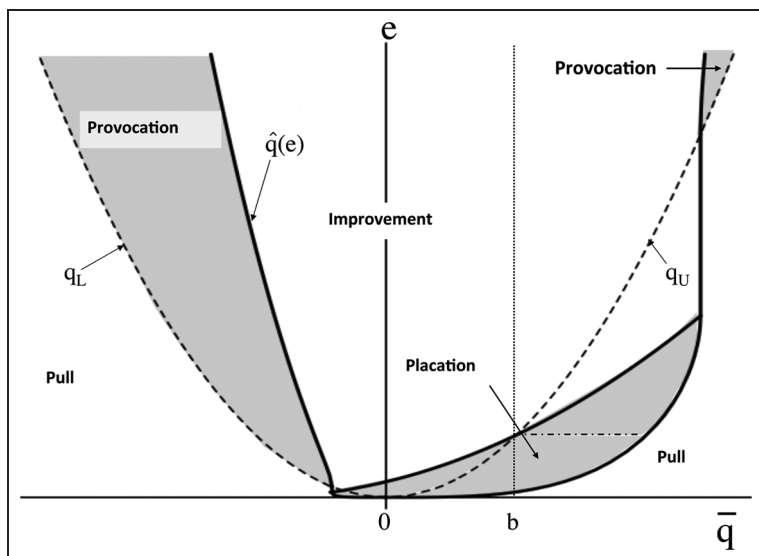


Figure 2. Equilibria in $\bar{q}$ - e space.

Note: In this graph $C_2(\cdot) = 0$.

in a policy on the bound of Player 2's inaction zone. The placation equilibrium regime is illustrated in Figure 4(a). The solid arrow indicates where Player 1 moves the policy.

The counterfactual in which the film industry decides not to self-regulate would represent a pull equilibrium. If censorship imposes low costs on the industry (bliss points are near each other), or if self-regulation would need to be very intense to placate (low e), or if self-regulation were very costly for the industry (high C_1), it would be in the industry's best interests to allow government to regulate. This seems to be quite evident in many industries in which no attempt to self-regulate is made in the face of imminent regulation from government. When Player 2's marginal costs are nonzero, both players move the policy toward their bliss points in a pull equilibrium. The lower the marginal costs of a player, the more that player will move the policy toward their bliss – the more “pull” that player exerts. In the special case when Player 2's marginal costs are zero, Player 1 takes no action, $m_1 = 0$, and Player 2 moves the policy to its bliss point, $m_2 = -\bar{q}$. This special case is depicted in Figure 4(b), where the dotted arrow indicates where Player 2 moves the policy.

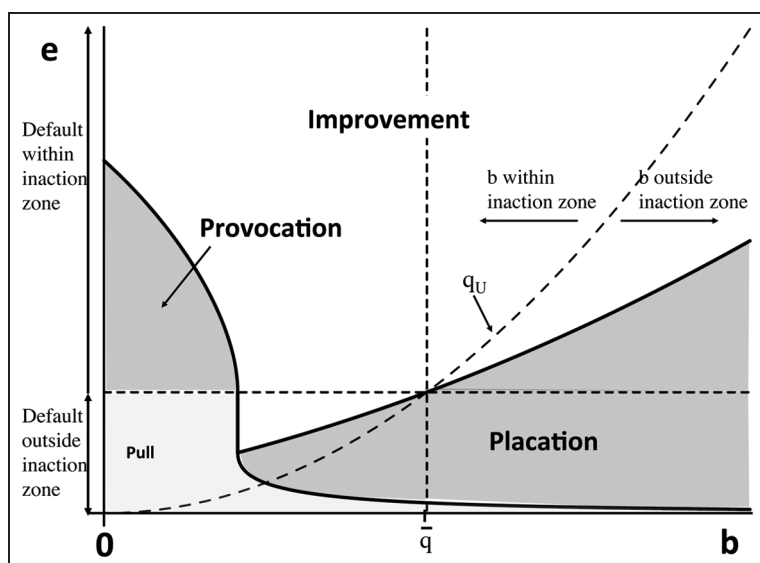


Figure 3. Equilibrium regions in b - e space.

Note: In this graph $C_2(\cdot) = 0$.

There is a second region of the parameter space that will induce this same equilibrium where the default policy begins within the inaction zone, $\bar{q} \in [q_L, q_U]$, Player 2 has high marginal costs, and Player 1 has large bias. Due to Player 2's large marginal costs, Player 1 is undeterred from moving the policy toward its bliss point. This can be seen in Figure 4(c).¹ This would be akin to the conflict of industry and regulation in a relatively lawless society where enforcement is weak. The industry resists and the government ineffectually regulates.

This application does not lend itself well to a provocation interpretation, and we can learn something by exploring why. It is difficult to think of a scenario in which a firm will make the policy intentionally worse in order to be regulated. This may be the case because the marginal costs for the industry to move the policy are near zero. For example, if the industry wants to lower the violent content in films, they can simply choose this violence level without paying any costs to lower the policy. Since they pay no costs, there is no incentive for them to free-ride on government's regulatory action.

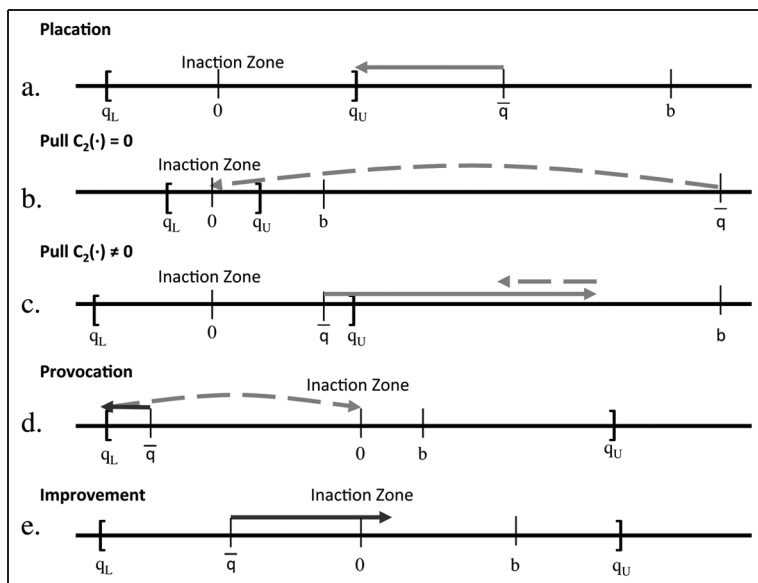


Figure 4. Placation, pull, provocation, and improvement regimes.

Note: The solid line represents Player 1's action and the dotted line represents Player 2's action.

Revolution: Provocation and improvement

The motivating example for the provocation logic is the ideology of revolutionaries. In the words of Thomas Jefferson, from *The Declaration of Independence*, "... All experience hath shown that mankind are more disposed to suffer, while evils are sufferable, than to right themselves by abolishing the forms to which they are accustomed." Several Latin American communist revolutionaries have espoused the philosophy that one must make evils beyond sufferable to spur the people to action. Perhaps the best known theorist on provocation was Carlos Marighela, who tried to start a revolution in Brazil (Wardlaw, 1989). In Marighela (1972):

The basic principle of revolutionary strategy in a country of permanent political crisis is to unleash, in urban and rural areas, a volume of revolutionary activities which will oblige the enemy to transform the country's political situation into a military one. Then discontent will spread to all social groups and the military will be held exclusively responsible for failures.

According to historian Richard Gott, Che Guevara believed in this revolutionary ideology as well (Marighela, 1972):

Guevara ... grasped the importance of an American invasion. In an apathetic continent, only direct engagement on the part of the Americans would stir up the necessary nationalism that could lead to a successful revolutionary war . . . Consequently, Guevara began a rural guerilla foco in Bolivia that was designed to bring American intervention and to spark off a continental war.

These quotes foreshadow the revolution in Nicaragua, where the communist Sandinistas ousted the Somoza family from power in 1979 (Nolan, 1984). In the time leading up to the revolution, the Somoza regime engaged in a repressive campaign against the peasants, including what is widely believed to be an extra-judicial assassination of a well-known critic of the regime, Pedro Joaquín Chamorro. The ranks of the Sandinistas rose as a consequence, and they soon had sufficient strength to gain control of the country.

This logic also closely parallels the strategy of al-Zarqawi, who aimed to start a civil war in Iraq following the US invasion. By attacking the Shia, al-Zarqawi expected a violent response which would exacerbate the condition of the Sunni (Hashim, 2006). Al-Zarqawi's strategy was to make the conditions for the Sunni sufficiently drastic that they would revolt against the Shia and the United States.

The model operationalizes these examples: the vanguard moves first and the populace follows. The populace has a fixed cost to revolting due to the collective action problem associated with coordination. Both the vanguard and the populace care about the allocation of goods in the economy, but the vanguard has a greater taste for equality. By making the distribution of resources for the populace worse to provoke them to action, this ultimately results in a better distribution for the populace, but at the cost of a revolution.

The provocation equilibrium occurs when the bias of Player 1 is not too large, and the default policy begins in the inaction zone of Player 2, $\bar{q} \in [q_L, q_U]$, and near one of the two bounds, q_L or q_U . Player 1 moves the policy in the opposite direction of its bliss in order to induce action on the part of Player 2, thereby free-riding on Player 2's efforts. The policy is made sufficiently undesirable such that it is in the interest of Player 2 to take action (although Player 2 is indifferent, Player 1 could move the policy an arbitrarily small amount to move the policy fully outside the inaction zone ensuring provocation). This is illustrated in Figure 4(d) in the case where the marginal cost of the populace is zero, $C_2(\cdot) = 0$. The solid

arrow shows where Player 1 moves the policy and the dotted arrow shows where Player 2 moves the policy. When Player 2's marginal costs are zero, Player 2 moves the policy directly to its bliss point.

The improvement equilibrium is the counterfactual in which would-be revolutionaries decide to act directly in the populace's best interest. One can think of this as the vanguard choosing to set up a soup kitchen rather than choosing to ignite a revolt. When would this happen? If the fixed costs to revolution are so high that drastic action is required to spur revolution, then improvement may be preferred. Alternatively if preferences between the vanguard and populace differ greatly (large b), or if the status quo is not very bad from the populace's perspective (default policy $\bar{q}$ near zero), setting up the soup kitchen may be a more desirable approach. This region is labeled "improvement" in Figure 1.

Swords or plowshares? When will Player 1 choose to provoke over improving the policy directly? Let us define $\hat{q}_E$ as the cutoff values for $\bar{q}$ at which Player 1, in equilibrium, is indifferent between moving the policy closer to its bliss point or moving the policy to the lower bound of the inaction zone and inducing a provocation equilibrium. Without loss of generality, we will focus on the provocation equilibrium involving the lower bound q_L . When $\bar{q} = \hat{q}_E$ both provocation and improvement are subgame perfect equilibria. This curve is depicted in $\bar{q}$ - b space in Figure 1, and it separates the provocation and improvement regions.

We already have a lower bound of $\bar{q}$ for which this provocation equilibrium is possible, and that is $\bar{q} = q_L$. Thus the region of $\bar{q}$ in which the subgame perfect equilibrium is a provocation is $[q_L, \hat{q}_E]$. How does this region change as a function of the divergence in preferences b between the two players?

There are two forces at play. As Player 1's preferences become more extreme relative to Player 2's, Player 1 cannot benefit much from free-riding on Player 2's actions, since its response will be moderate. Revolutionaries cannot start a communist revolution if the populace is fond of inequality. This first force makes provocation less attractive to Player 1. On the other hand, as Player 1 becomes more extreme, the status quo appears less favorable. In other words, as the revolutionaries become more extreme the current level of inequality is disliked more. With this second force, greater extremism makes improvement less desirable relative to provocation for Player 1. When b is very large, there will be no default values that lead to provocation. Player 2's response will not be sufficient to make provocation a desirable choice for Player 1.

Proposition 1. There exists a bias, $\hat{b}$, such that $\forall b > \hat{b}$, provocation no longer exists as an equilibrium for any $\bar{q}$.

Player 1 will provoke Player 2 only if their preferences are closely aligned. For example, a group that opposes the ruling party and has fairly similar interests to the populace may wish to start a rebellion as in the examples presented in the introduction. However, when the group is very extreme (e.g. terrorist organizations or extreme ideologues), it seems implausible that a group would provoke a population that disagrees with their vision. In such cases it is more likely that the extreme group will explore alternative roads to power, such as through capturing the seats of power.

A provocation equilibrium is an undesirable outcome for Player 2. A provocation equilibrium requires that the default policy begins within the inaction zone and that Player 1 pushes it out of the zone. Player 2 always prefers Player 1 to leave the policy within the inaction zone. Thus, there are regions of the default space for which Player 2 is *better off* when Player 1 is *more* extreme.

For proof of Proposition 1, see Appendix 1.

Applications

Frontier conflict

The interaction between the revolutionaries and the populace is an interaction of provocation. The same structure can be seen with the interaction between settlers on a frontier and the military at a base in the homeland. History is rife with conflicts between settlers and natives. In a PBS Frontline documentary, an Israeli settler living in close proximity to Palestinians says (Setton, 2005):

And it's only a matter of time until the war, with God's help, will begin, and it will begin with us. And in the end, we'll win. We'll inherit the land and expel the people who are in it.

The settler wants to provoke a conflict with the Palestinians. In this sort of conflict, the military must decide how much assistance they will grant to the settlers. The settlers are Player 1 and the military is Player 2.

The settlers (Player 1) may want to initiate a conflict with the natives in order to provoke military assistance. In this sort of conflict, the military (Player 2) must decide how much assistance they will grant the settlers. This is the same interaction as the one sometimes (mis)referred to as the

“moral hazard” of third-party intervention (Kuperman, 2008; Rauchhaus, 2005).

In a general sense, settlers can choose to pursue relations with the natives in either a cordial manner or an aggressive manner. In our simplified abstraction, settlers care only about access to land and its resources. Map the settlers’ access to the land onto a single policy dimension. A low value of $\bar{q}$ can be interpreted as natives who grant limited access to the settlers, whereas a high value of $\bar{q}$ can be interpreted as natives who grant wide access to the settlers. When Player 1 moves the policy higher toward greater access, this can be interpreted as having a cordial relation with the natives, with a marginal cost of giving tribute. Moving the policy leftward, toward lesser access, can be interpreted as engaging in aggressive relations, with a marginal cost to instigating skirmishes.

The military wants the settlers to have a high level of access to the land, but not as high as the settlers want for themselves. The military can attack the natives to improve the settlers’ access but doing so incurs a fixed cost. It requires logistics, recruitment, reconnaissance, etc. Thus the military will only get involved if the settlers’ access to resources is very poor. A pull equilibrium occurs when the initial access is poor and the military intervenes to help the settlers by fighting off the natives. A provocation equilibrium occurs when the initial access is poor, and the settlers challenge the natives in order to draw the military into the conflict. An improvement equilibrium occurs when the initial access is good, and the settlers have good relations with the natives.

The military faces a large frontier full of many settlements interacting with different native groups. Each interaction is drawn from some known distribution. In some cases the settlers might have good access (high $\bar{q}$), and in other cases the access might be poor (low $\bar{q}$). What if the military could control the size of the fixed costs (e) to aiding the settlers? The military might construct bases closer to the frontier to decrease the fixed costs of intervention, or they could choose to construct bases far from the frontier in order to keep the fixed costs high.

What is the optimal fixed cost for Player 2, the military? In the provocation equilibrium Player 1 engages in Pareto-damaging behavior that lowers Player 2’s utility, and in a pull equilibrium Player 1 may shade down its action and free-ride on Player 2’s action. In some cases, it will be in Player 2’s interest to have larger fixed costs in order to deter Player 1 from provoking and then free-riding on Player 2’s actions. This serves as a credible commitment to Player 1 that Player 2 will take no action. Figure 5 shows a graph of Player 2’s utility for some fixed $\bar{q} < 0$ and a fixed b . As the fixed

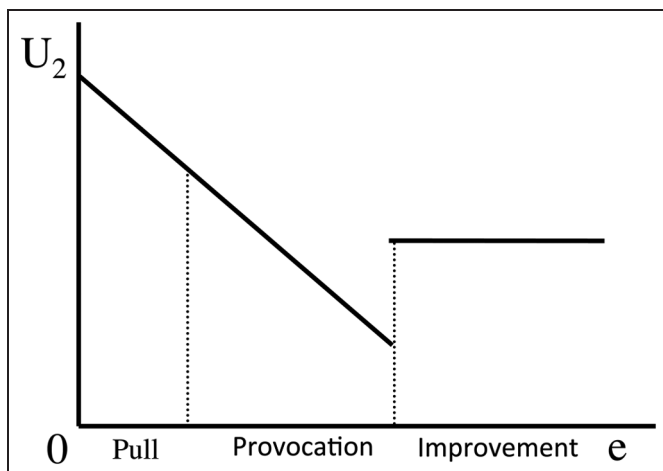


Figure 5. An example of Player 2's utility as a function of the fixed cost.

Note: In this graph $C_2(\cdot) = 0$.

costs for Player 2 increase, the equilibrium regime changes from pull, to provocation, to improvement. Initially, increasing the fixed costs results in a loss. Player 2 pays these fixed costs in a pull equilibrium. As e increases, Player 1 decides instead to improve the situation directly. Once the equilibrium is an improvement regime, increasing the fixed costs further has no effect on the equilibrium or the utility of either player.

The military will choose the location of the base depending on the initial distribution of settler resource access. If the distribution is heavily weighted toward defaults with very low access, the military will place the base near the frontier in order to easily intervene. If the distribution is weighted toward poor to mediocre access, then the military will place the base far from the frontier in order to induce the settlers to extend relations to the natives.

Expert supervision

Consider a firm with a hierarchical structure. The expert supervisor oversees many subordinates, each of whom works on his own project. A subordinate is Player 1 and the supervisor is Player 2. There is an inherent tradeoff between the quality of the final outcome of a project and the time

spent on a project. For instance, q may be the time and quality per project. The supervisor may have a stronger preference for speed while the workers may have a stronger preference for quality. Projects differ in the amount of time it requires to reach a certain level of quality. This translates to a distribution of $\bar{q}$. For example, a low $\bar{q}$ can be interpreted as a project that requires a lot of time to complete at a decent level of quality, whereas a high $\bar{q}$ can be interpreted as a project that requires little time to complete at a decent level of quality.

The supervisor may assist the subordinates on their projects. If the supervisor has expertise that the subordinates lack, we would expect that her marginal costs to working on a project are much lower than the subordinates' marginal costs. However, we may also expect that the supervisor has fixed costs to aiding a subordinate. It may take a large amount of time for the supervisor to familiarize herself with the necessary details of a project in order to be of any assistance. As a particular instance, consider the advisor–advisee relationship in academia. The advisor may be much more efficient in solving roadblocks on a research project, but only the advisee is familiar with the details of the roadblock and how it is embedded in the project more generally.

When a project is very difficult ($\bar{q}$ is very low), it is best for the supervisor to aid the subordinate (pull equilibrium), and if a project is easy (high $\bar{q}$) the subordinate finishes the project solo (improvement equilibrium). However, if the project is of medium difficulty (intermediate $\bar{q}$), the subordinate may neglect necessary maintenance in order to allow the quality to become sufficiently low to elicit the aid of the supervisor (provocation equilibrium). If many projects fall under this category, the supervisor would wish to tie her own hands in order to deter this Pareto-damaging behavior.²

Suppose upper management can choose what kind of supervisor to hire and their preferences are the same as the supervisor's. Upper management can promote someone from within the organization who has detailed knowledge of the projects (low e), or they can hire an outsider who is competent but has little specific knowledge of any projects (high e). If there are many projects of intermediate difficulty, they might hire the outsider in order to encourage subordinates to be more self-sufficient.

Self-control issues

Often individuals want their future selves to take a particular action, but when the time comes lack of self-control causes them to take another action. Consider a consumer that is impatient and has present-biased

preferences (O'Donoghue and Rabin, 1999). This present-biased agent faces the challenge of adhering to a diet. On the first period, this individual purchases (or packs) his food, and in the second period he consumes. Ideally he would like to consume only the healthiest of foods in the second period. However, he is sophisticated and aware that if he does not provide a dessert for himself, he will renege on his diet in the second period. We assume there is a fixed cost of time and effort associated with the shopping required to change the planned diet. If the consumer buys a small dessert this may be sufficient to placate his period-2 self. If his meal does not include a dessert, his cravings may overwhelm him and he will walk down the block to the local convenience store and buy his favorite candy bar. This is a placation equilibrium analogous to the regulation example and illustrated in Figure 4(a).

We can also consider situations in which a present-biased agent may provoke his future self. Consider a worker who aims to have an ideal balance between work and leisure over the course of his workday and weekend. He is Player 1 on the workday, and he is Player 2 on the weekend. He only cares about the displeasure of effort on the day that he exerts the effort. There is also a work–leisure tradeoff. Map the preferences for this tradeoff onto the real number line. So a bliss point represents the ideal balance between work and leisure. We can model $\bar{q}$ as an exogenous shifter of the optimal amount of work. Higher exogenous incentives for work (low $\bar{q}$) mean that more effort must be exerted to reach the same work–leisure ideal as when there are lower stakes (higher $\bar{q}$). The agent has a fixed cost to work on the weekend: he must go out of his way to the place of employment, whereas on the workday this is a sunk cost.

On projects with high stakes (low $\bar{q}$) the agent knows he will go to the office in both periods, so he shades down his effort on the weekday (pull equilibrium). On projects with low stakes (high $\bar{q}$), the agent works only on the weekday (improvement equilibrium). On weeks of moderate stakes (intermediate $\bar{q}$), the period-1 agent knows that he cannot get his future self to commit to working on the weekend even though he would like himself to do so. Thus on the weekday the agent greatly shades down his work (i.e. procrastination) because it is the only way to provoke himself to work on the weekend (provocation equilibrium).

Over the course of the agent's job tenure, the incentives of the projects will be drawn from a distribution. Before the agent starts his job, he can choose how far from work to live and consequently determine his fixed costs (e) to working on the weekend. He must balance out the costs of commuting when working on high-stakes projects against having a good

commitment device to not working on the weekend when there are moderate stakes, thereby deterring provocation (i.e. preventing procrastination). When the typical project has moderate stakes, the agent may choose to increase the separation between work and leisure to provide a commitment device against weekday procrastination. These costly barriers make the time at work more productive.

Discussion and extensions

Assumptions

A key assumption in the model is that Player 2 does not have the power to directly penalize Player 1. In the self-regulation example, the government could change the level of “iniquity” in the media, but they cannot directly penalize the industry in this model. This is a realistic assumption in some cases. The government cannot legally prosecute the industry before they legislate regulations, hence they cannot directly punish firms. In the context of revolution, the populace may wish to punish the vanguard if the vanguard makes the situation worse, but they are unorganized and unarmed. If they had the organization and arms to punish the vanguard, then it seems likely that they would have the power to directly change the policy by pressuring the oppressive government (who is not modeled as a player). Likewise, the oppressive regime may want to directly punish the vanguard, but it may not have the means to do so without oppressing the whole populace. It is plausible that the vanguard can hide amongst the people and wage protracted guerilla warfare using civilians as a shield. If the oppressive regime can single out the vanguard and punish them without affecting the rest of the population, then the model ceases to apply.

Another assumption that should be addressed is the origin of the default policy, $\bar{q}$. In many political economy models the default policy is considered exogenous. This paper takes the same approach. Our interpretation is that the location $\bar{q}$ was determined as the outcome of some other unrelated interaction. The fact that movies included sex and violence in the 1930s may be the result of many socioeconomic factors. The fact that the populace is oppressed may be the outcome of a long, complicated history of the state. Both of these causes are beyond the scope of the model.

Many firms

Previously, a whole industry was modeled as a single agent. In many circumstances one would expect that the first mover may comprise a group of

independent agents. For example, some firms may choose to self-regulate whereas others may not agree to the sanctions. In this section and the next the model is extended to allow for multiple first movers. An industry is modeled as an arbitrary number of firms, and a revolutionary vanguard is modeled as an arbitrary number of factions.

In both the placation and provocation equilibria, the first mover moves a policy away from its bliss point. In this section we find that, when the number of first movers is large, these equilibria actually become more favorable to the first movers relative to the pull and improvement equilibria. The intuition is that the pull and improvement equilibrium regimes exhibit free-riding effects between the first movers, while the placation and provocation equilibrium regimes involve a “provision point” that eliminates any free-riding. With the provision point all the first movers are pivotal.

Suppose the industry is composed of N identical firms and the government is considering censorship. We construct the model to keep the total costs and benefits of the industry invariant to the number of firms. Thus as we explore the comparative statics with respect to N keeping constant the benefits and costs of the industry, let $M = \sum_{k=1}^N m_k$ be the sum of the actions of the first movers and $M_{-i} = M - m_i$ be the sum of the actions of the first movers excluding the action of player i . The new utility functions are:

$$U_i(m_i, m_{-i}) = \frac{1}{N} f_F(\bar{q} + M + m_G) - m_i * c_F \quad (3)$$

$$U_G(m_{-G}, m_G) = f_G(\bar{q} + M + m_G) - C_G(m_G) - e * 1\{m_G \neq 0\} \quad (4)$$

Notice that $\sum_{i=1}^N U_i(m_i, m_{-i})$ is equal to the original utility function in equation (1), except now individual firm cost functions have constant marginal cost. This cost function is chosen because for a given M total industry costs are invariant over the distribution of actions m_i across the first movers. Thus it allows for a more direct comparative static as N varies. For the following analysis, assume that marginal costs c_F are sufficiently small so that some nonzero action on the part of a firm is optimal when $N=1$. Otherwise the firm will do nothing and the analysis ends there. Furthermore, assume the government's marginal costs are increasing $C_G''(m_G) > 0$.

Previously, we compared the placation equilibrium to the pull equilibrium. There exists a single value of $\bar{q}$ for which both are equilibria (as b varies this single value becomes a curve which is shown separating the pull region from the placation region in Figure 1). By introducing multiple

firms in this extended model, coordination becomes an issue. Equilibria are no longer generically unique. There may be a region where placation and pull are equilibria depending on how firms coordinate. As N increases, all equilibria become more robust, because as the number of firms increases (and the total utility of the industry is kept constant), each firm has a smaller influence on the market. A small firm will incur exorbitant costs relative to its size if it deviates in an attempt to substantially change the final policy.

First, we consider the multiple-player equivalent of the placation equilibrium. Total action by the first movers moves the default policy to the upper bound of the inaction zone, q_U . Let us consider only a symmetric placation equilibrium. This is the placation equilibrium that can be induced for the largest space of $\bar{q}$.³ Define the region $[q_U, \hat{q}_P(N))$ as the region of $\bar{q}$ where a symmetric placation equilibrium can be coordinated upon.

Now consider the multiple-player equivalent of the pull equilibrium. Define the region $[\hat{q}_G(N), \infty)$ as the region of $\bar{q}$ where a pull equilibrium can be coordinated upon. We wish to know how $\hat{q}_P(N)$ and $\hat{q}_G(N)$ change as a function of N . Notice that when there is one firm, $\hat{q}_P(1) = \hat{q}_G(1)$.

Lemma 1. *There exists an N' such that for all $N > N'$, the optimal firm action in a pull equilibrium is $m_i^*(N) = 0$, and $\hat{q}_G(N) < \hat{q}_P(N)$.*

Lemma 1 states that when the industry is composed of many firms, the region of $\bar{q}$ that induces a placation equilibrium, $[q_U, \hat{q}_P]$, overlaps with the region of $\bar{q}$ that induces a pull equilibrium, $[\hat{q}_G, \infty)$. So firms can coordinate on either equilibrium. The lemma also says that the optimal firm action in a pull equilibrium decreases to zero. The intuition here is that there is a free-rider effect. When there are many firms, each firm will do less. When there are very many firms, each firm will do nothing. This comes from decreasing marginal benefit to acting as N increases, but constant marginal cost.

Suppose $\bar{q} \in [\hat{q}_G, \hat{q}_P]$ and the firms can coordinate on a subgame perfect equilibrium. Which equilibrium regime would the firms choose? Define $U_{plac}(N; \bar{q}) \equiv \sum_{i=1}^N U_i(m_i, m_{-i}; \bar{q})$ as the total utility of the firms in a symmetric placation equilibrium. Define $U_{grav}(N; \bar{q}) \equiv \sum_{i=1}^N U_i(m_i, m_{-i}; \bar{q})$ as the total utility of the firms in a pull equilibrium.

For proof of Lemma 1, see Appendix 1.

Proposition 2. *If $N_1 < N_2$ and $\bar{q} \in [q_U, \hat{q}_P(N_1)]$ then $U_{plac}(N_1; \bar{q}) = U_{plac}(N_2; \bar{q})$. Furthermore, there exists an N' such that for all $N > N'$ and $\bar{q} \in [\hat{q}_G(1), \infty)$, $U_{grav}(N; \bar{q}) < U_{grav}(1; \bar{q})$.*

Proposition 2 states that equilibrium utility in a symmetric placation equilibrium is invariant to N while the equilibrium utility in a pull equilibrium decreases when N is large. Specifically we can say that if $N = 1$ and $\bar{q} = \hat{q}_P(1) = \hat{q}_G(1)$, then as N increases, the symmetric placation equilibrium becomes preferred by all firms to the pull equilibrium. Note also that the government does better in this equilibrium as well. The firms and the government would want coordination on a placation equilibrium.

The intuition is that there are positive externalities when a firm moves the policy in a pull equilibrium. If a firm fights against regulation, all the other firms benefit without paying the cost. Each firm only internalizes its individual benefit. However, in a placation equilibrium, all players are pivotal and thus if any one firm shirks, the policy will not reach the edge of the inaction zone and consequently the government will regulate. So if any firm fails to self-censor they all get regulated. This is like a “provision point” in the public goods literature (Isaac et al., 1989; Marwell and Ames, 1980).

The public goods analogy extends a bit further. One can interpret the policy as the provision of a public good where $\bar{q}$ is the initial exogenous provision. The policy is a non-rivalrous public good in the sense that a firm’s consumption of the policy does not detract from any other firm’s consumption. The interaction between the firms is like a voluntary contribution mechanism (VCM) with a provision point. The firm’s contribution is the degree to which it moves the policy.

In the pull equilibrium as N increases, the firms internalize less of the benefit of their actions, and hence the positive externality problem worsens. Since the provision point in the placation equilibrium solves the positive externality problem, as N increases, the industry has a larger incentive to coordinate on a placation strategy. This is Pareto optimal since no firm can do better without the government or another firm doing worse, and the government can do no better without a firm doing worse.⁴

For proof of Proposition 2, see Appendix 1.

Split revolutionary factions

How does factionalization affect the revolutionaries’ incentives to provoke? Let there be N identical revolutionary factions, all have the same bliss point, and each can act independently. They have utility functions given by equation (3). They face a single populace that has a utility function given by equation (4). The default policy, $\bar{q}$, begins within the inaction zone and it is assumed without loss of generality that $b > \bar{q}$. This section

shows that the region of $\bar{q}$ that results in both improvement and provocation equilibrium regimes expands as N increases. In addition, the relative equilibrium utility of the factions in a symmetric provocation equilibrium increases relative to the equilibrium utility of an improvement equilibrium.

As in the previous section, when $N > 1$ the equilibrium requires coordination between the factions. There may be regions of the default space ($\bar{q}$) where both improvement and provocation are equilibria. Define $[q_L, \hat{q}_R(N)]$ as the region of $\bar{q}$ where a symmetric provocation is an equilibrium and define $[\hat{q}_{sl}(N), \hat{q}_{su}(N)]$ as the region of $\bar{q}$ where improvement is an equilibrium. Notice that $\hat{q}_R(1) = \hat{q}_{sl}(1)$.

Lemma 2. *The total optimal first-mover action in an improvement equilibrium, $N * m_i^*(N)$, is monotonically decreasing in N . There exists an N' such that for all $N > N'$, $m_i^*(N) = 0$ and $\hat{q}_{sl}(N) < \hat{q}_R(N)$.*

Lemma 2 says that as the vanguard splinters into more factions, the region of $\bar{q}$ that results in a symmetric provocation equilibrium and the region of $\bar{q}$ that results in an improvement equilibrium expands as well. The intuition is the same as in the previous subsection. If we keep the total vanguard utility constant, splintering the vanguard into more factions gives each faction less power to affect the final outcome. When N is large, $\bar{q}$ must be extreme for a single deviation to be profitable. Lemma 2 also states that the total optimal action in an improvement equilibrium will reduce to zero.

Suppose $\bar{q} \in [\hat{q}_{sl}(N), \hat{q}_R(N)]$ and the firms can coordinate on a subgame perfect equilibrium. Which equilibrium regime would the firms choose?

Define $U_{rev}(N; \bar{q}) \equiv \sum_{i=1}^N U_i(m_i, m_{-i}; \bar{q})$ as the total utility of the factions in a symmetric provocation equilibrium. Define $U_{improve}(N; \bar{q}) \equiv \sum_{i=1}^N U_i(m_i, m_{-i}; \bar{q})$ as the total utility of the factions in an improvement equilibrium.

For proof of Lemma 2, see Appendix 1.

Proposition 3. *If $N_1 < N_2$ and $\bar{q} \in [q_L, \hat{q}_R(N_1)]$ then $U_{rev}(N_1; \bar{q}) = U_{rev}(N_2; \bar{q})$. Furthermore, there exists an N' such that for all $N > N'$ and $\bar{q} \in [\hat{q}_{sl}(1), \hat{q}_{su}(1)]$, $U_{soup}(N; \bar{q}) < U_{grav}(1; \bar{q})$.*

Proposition 3 states that equilibrium utility in a symmetric provocation equilibrium is constant while the utility in an improvement equilibrium is lower when N is large. Specifically, we can say that if $N=1$ and $\bar{q} = \hat{q}_R(1) = \hat{q}_{sl}(1)$, then increasing N causes all factions to prefer the provocation equilibrium. The factions would prefer to coordinate on

provocation. However, this provocation equilibrium is not Pareto optimal because a provocation equilibrium always makes the populace (Player 2) worse off than in an improvement equilibrium.

The intuition for this result is that the edge of the inaction zone (q_L) provides a provision point for the factions in the provocation by extremists' equilibrium. There is no free-riding since all factions are pivotal. A single deviation would result in a failed provocation. In an improvement equilibrium, however, the free-riding problem gets worse as N increases. So a factionalized vanguard will free-ride on each other's efforts to improve the political situation, and this makes them collectively ineffectual. The promise of revolution provides incentives that eliminate shirking.

In both this subsection and in the previous subsection (on multiple firms), the analysis is predicated on the assumption that there is perfect coordination. There also exist mixed strategy equilibria in which all N firms or factions mix their amount of effort. In this type of equilibrium, there may be a high probability that the effort exerted by the firms or factions falls short of the provision point: placation and provocation may rarely occur.

In this game, as with many games that have a Stag-Hunt-like structure,⁵ one may worry that strategic uncertainty about others play could unravel the placation equilibrium. However, there are also reasons to believe that firms or factions may be able to solve this strategic uncertainty problem. Dietz et al. (2003) review ways in which a collection of individuals succeed in managing common resources. They find that even without provision points, observability and use of sanctions can help to induce optimal resource use by the individuals. Industries in which firm behavior is publicly observable and in which business associations can provide internal incentives to their constituent members are likely to be more effective in coordinating behavior on the placation equilibrium. Likewise, political movements in which factions can credibly communicate and reward each other with side compensations are likely to be more effective in coordinating behavior on the provocation equilibrium.

For proof of Proposition 3, see Appendix 1.

Asymmetric information

Firms may not know exactly the government's fixed costs to regulating, and revolutionaries may not know exactly the populace's fixed costs to revolt. Asymmetric information adds an additional strategic layer to the model. Suppose that Player 2 knows its fixed cost but all the first movers only know the distribution from which e is drawn. Does this eliminate the

provision point? Not necessarily: if e is distributed discretely then each mass point could potentially become a provision point. However, if e is distributed continuously, this eliminates the provision point. Suppose e is continuously distributed from e_1 to e_2 with $e_1, e_2 \in \mathbb{R}_{++}$ and $e_1 < e_2$. This implies that the upper bound of the inaction zone will also be continuously distributed. Let the support of the upper bound be $[q_U^1, q_U^2]$.

First, consider the case where $N = 1$. Perhaps the film industry knows that it is costly for the government to establish a censorship bureau, but it does not know how costly. Thus the location of the bounds of the inaction zone will be uncertain for Player 1. This uncertainty does not fundamentally change the placation logic. Player 1 will not know the exact location of q_U , but Player 1 can still move the policy toward Player 2's bliss point and away from its own bliss point such that the probability that the policy is within the inaction zone is sufficiently high. Thus the qualitative aspect of placation remains, and Player 1 moves the policy toward Player 2's bliss point in order to deter action. The logic is exactly the same for the provocation equilibrium.

Asymmetric information will interact with the number of first movers. Consider the case when $N = 1$ and the optimal action for Player 1 is a placation or provocation equilibrium in which the optimal action is to move the policy $m_1 = \tilde{q} - \bar{q}$ where $\tilde{q} \in [q_U^1, q_U^2]$. Player 1 moves the policy such that the probability of placating (or provoking) Player 2 is sufficiently high and moving the policy any further is not worth the marginal cost. Now consider the equivalent placation (or provocation) strategy in which $N > 1$ first movers are collectively moving the policy $\frac{N-1}{N}(\tilde{q} - \bar{q})$, and consider a deviation by the N th player. There no longer exists a provision point where utility discontinuously increases when the provision point is reached. Hence the N th player's optimal action will satisfy its first-order condition, and since $\frac{N-1}{N}$ of the benefit of its action is externalized to the other players, the firm will shade down its action. As N increases the free-rider problem becomes more severe.

In light of asymmetric information, the results from 'Many firms' and 'Split revolutionary factions' sections can be seen as part of the incentive structure for the fixed-cost player to reveal or obscure the fixed cost. Suppose that at a pre-stage of the game before any player knows the fixed cost, Player 2 may publicly and credibly reveal the fixed cost or keep it occluded. If the fixed cost is revealed, the game proceeds as described in these sections. If the fixed cost is occluded, all players know the continuous distribution from which the fixed cost is drawn but not the fixed cost. The first movers make their move, then Player 2 discovers the fixed cost, and then moves.

As shown in Proposition 2, when N is large, and the fixed cost is publicly known, the region of the parameter space that results in a placation equilibrium is larger relative to when $N = 1$. However, when the fixed cost is uncertain to the first movers, the expansion of the placation equilibrium in the parameter space does not occur. Thus in order to induce a placation equilibrium, Player 2 is more likely to reveal the fixed cost when N is large relative to when $N = 1$. In other words, as N increases, Player 2's incentive to reveal information about its fixed cost increases as well.

A way to think of this is that the government is more likely to reveal the threshold of its tolerance for violent films when there are many firms. When there is only one firm, the government may wish to occlude the threshold of its tolerance in order to induce the single firm to overshoot in its placation. However, when there are many firms, occlusion will simply result in costly regulation – the firms cannot self-regulate due to the free-rider problem. Thus, to provide a provision point to ensure self-regulation, the government reveals its threshold.

The same logic applies to the populace's incentive to reveal or occlude their fixed costs when faced with a factionalized vanguard, but the result is the reverse. In order to prevent a provocation equilibrium, the fixed-cost player will occlude the fixed cost for a larger parameter space when N is large relative to when $N = 1$. In other words, as N increases, Player 2's incentive to occlude information about its fixed cost increases as well.

A way to think of this is that when the revolutionaries are a tight one-unit organization, the populace does not hide its threshold of tolerance, because the revolutionaries will provoke regardless. However, when the revolutionaries are factionalized, the populace can hide their willingness to fight to dissuade provocation. This introduces free-riding problems should the factions attempt a revolution. As a consequence, the factions are induced to improve the political situation directly.

Multi-dimensional interests

In some cases, the conflict of interest between the first movers and Player 2 may be over multiple dimensions. For example, the firms and the government may care about both the safety of the product and also protectionism. The government prefers more safety than do the firms, and the firms prefer more protectionism than does the government. Suppose we modify the model such that the players have preferences over K policy dimensions rather than just a single dimension. Assume that preferences remain single-peaked, and that the marginal cost of moving the policy is a function

of the distance moved. Observe that Player 2 will still have an inaction zone around its bliss point. If Player 2's preferences are symmetric (i.e. $f_2(q) = f_2(-q)$), then the inaction zone will take the form of an n -dimensional ellipsoid. In any case, the logic of placation and provocation follows naturally. If the policy, $\bar{q} + m_1$, is near Player 2's bliss point on Player 2's turn, then Player 2 will take no action, just as in our basic model. The intuition remains the same.

Repetition

The vanguard and the populace may have a repeated strategic relationship. A revolution may be followed by additional attempts to influence the political climate by the vanguard. Suppose the game is repeated an infinite number of times where future utility is exponentially discounted, and the equilibrium policy from the previous round, $q^{t-1} = \bar{q}^{t-1} + m_1^{t-1} + m_2^{t-1}$, becomes the default policy in the current round $\bar{q}^t$. For games like this, it is usually helpful to restrict attention to Markov strategies, where players condition their strategy solely on the state $\bar{q}^t$, and not on the history. Assume that a stage-game outcome that places the policy closer to Player i 's bliss point in period t makes Player i better off in the infinitely repeated game. As long as this is true, Player 2's fixed costs will create an inaction zone. If Player 2 takes action, it produces benefit not only in the current period but also in all subsequent periods. Nonetheless, for policies that are very close to Player 2's bliss point, the net-present value of taking action is insufficient to outweigh the fixed costs. This presence of the inaction zone allows the logic of placation and provocation to follow.

Conclusion

This paper presents a simple lens in which to view many strategic situations. The advantage of simplicity is that it allows for wide application across disparate phenomena. The sacrifice is depth for any particular application. The model is not designed for thorough analysis of any one application but rather to uncover fundamental commonalities across applications.

The model leads to four main predictions. First, agents that have fixed costs may be better off when their preferences are misaligned with would-be provocateurs. Since the motivation to provoke is based on the free-riding of Player 2's action, misaligned preferences yield an undesirable action from the first movers' perspectives. Likewise, if Player 2 can commit to a large fixed cost or convince the first movers that it has a large

fixed cost, this too can discourage provocation because it makes provocation more costly. Free-riding becomes an issue amongst the various first movers when comparing provocation to direct improvement, or when comparing placation to directly moving the policy toward one's interests (i.e. pull equilibrium). The model predicts that under perfect information, provocation and placation are more likely when there are many first movers. These two regimes provide a provision point that eliminates free-riding. Consequently, this gives Player 2 incentives for occluding or revealing its fixed cost and affecting the presence of the provision point. The number of first movers increases Player 2's incentive to reveal its fixed cost in order to aid in placation. In contrast, the number of first movers increases Player 2's incentive to occlude its fixed cost in order to deter provocation.

Acknowledgements

Earlier versions of this draft were circulated under the title "Threshold Activism." I am grateful to John Morgan, Botond Köszegi, Matthew Rabin, Gerard Roland, Ernesto Dal Bó, Matthew R. Levy, Claudia Sitgraves, Michael Katz, and Robert Powell for valuable comments. I would like to thank my anonymous referees for very insightful and helpful comments. I also thank seminar participants at U.C. Berkeley's Comparative Economics Seminar and NBER's Political Economics Student Conference. All errors are mine alone.

Funding

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Notes

1. This region of $\bar{q}$ cannot result in pull if Player 2's marginal costs are zero. This equilibrium is not depicted in Figure 1 because the functions used to generate the graph have $C_2(\cdot) = 0$.
2. In some situations it may be plausible for these strategic problems to be contracted away. However, in other situations, the costs to drafting, monitoring, and enforcing the contract may be more costly than no contract.
3. Asymmetric placation equilibria can only be induced from a strict subset of this interval.
4. In the special case where the marginal costs of the government are zero ($C_G(m_g) = 0$) this result no longer holds. In a pull equilibrium the government moves the policy all the way to their bliss and the firms take no action. The resulting utility is invariant to N . Thus as N increases, there is no effect on the region of $\bar{q}$ that results in a pull equilibrium.

5. A Stag Hunt is a 2×2 normal-form game with two pure-strategy Pareto-ranked Nash equilibria, and in which the Pareto inferior Nash equilibrium is risk-dominant.
6. For (A.3) and (A.4), both sides have been multiplied by N .
7. For (A.5) and (A.6), both sides have been multiplied by N .

References

- Acemoglu D and Robinson JA (2001) A theory of political transitions. *American Economic Review* 91(4): 938–963. Available at: <http://www.jstor.org/stable/2677820>
- DeMarzo PM, Fishman MJ and Hagerty KM (2005) Self-regulation and government oversight. *Review of Economic Studies* 72(3): 687–706.
- Dietz T, Ostrom E and Stern PC (2003) The struggle to govern the commons. *Science* 302(5652): 1907–1912.
- Ellis JC and Wexman VW (2002) *A History of Film*. Boston, Massachusetts: Allyn & Bacon.
- Grossman HI (1991) A general equilibrium model of insurrections. *American Economic Review* 81(4): 912–921.
- Hashim A (2006) *Insurgency and Counter Insurgency in Iraq*. Ithaca, New York: Cornell University Press.
- Isaac RM, Schmidt D and Walker JM (1989) The assurance problem in a laboratory market. *Public Choice* 62: 217–236.
- Kuperman AJ (2008) The moral hazard of humanitarian intervention: Lessons from the Balkans. *International Studies Quarterly* 52: 49–80.
- Marighela C (1972) *For the Liberation of Brazil* (trans. Butt J and Sheed R, with an introduction by Gott R). Harmondsworth, England: Penguin Books.
- Marwell G and Ames RE (1980) Experiments on the provision of public goods. ii. provision points, stakes, experience, and the free-rider problem. *American Journal of Sociology* 85(4): 926–937.
- Mast G and Kavin BF (1992) *A Short History of the Movies*. 6th ed. Boston, Massachusetts: Allyn & Bacon.
- Maxwell JW, Lyon TP and Hackett SC (2000) Self-regulation and social welfare: The political economy of corporate environmentalism. *Journal of Law & Economics* 43(2): 583–617.
- Nolan D (1984) The ideology of the Sandinistas and the Nicaraguan revolution. Technical report, Institute of Inter American Studies, Graduate School of International Studies, University of Miami, Coral Gables, FL.
- O'Donoghue T and Rabin M (1999) Doing it now or later. *American Economic Review* 89(1): 103–124.
- Rauchhaus RW (2005) Conflict management and the misapplication of moral hazard theory. *Ethnopolitics* 4(2): 215–224.
- Roemer JE (1985) Rationalizing revolutionary ideology. *Econometrica* 53(1): 85–108.

- Setton D (2005) Israel's next war? (transcript). In: *PBS's Frontline*. Available at: <http://www.pbs.org/wgbh/pages/frontline/shows/israel/etc/script.html>
- Stango V (2003) Strategic responses to regulatory threat in the credit card market. *Journal of Law & Economics* 46: 427–452.
- Stefanadis C (2003) Self-regulation, innovation, and the financial industry. *Journal of Regulatory Economics* 23(1): 5–25.
- Tullock G (1971) The paradox of revolution. *Public Choice* 11(1): 89–99.
- Volden C and Wiseman AE (2008) A theory of government regulation and self-regulation with the specter of nonmarket threats. Unpublished manuscript.
- Wardlaw (1989) *Political Terrorism: Theory, Tactics and Counter-Measures*. New York: Cambridge University Press.

Appendix I

Proof of Proposition 1

The utility of Player 1 in a provocation equilibrium in which Player 1 provokes by moving the policy to the lower bound of the inaction zone is

$$U_1^{provocation} \equiv f_1(q_L + m_2^*) + C_1(\bar{q} - q_L) \quad (\text{A.1})$$

Consider the utility of Player 1 if he were to move the policy to the upper bound of the inaction zone

$$\widetilde{U}_1 \equiv f_1(q_U) + C_1(q_U - \bar{q}) \quad (\text{A.2})$$

First notice that the costs are invariant to b . The difference in felicities, given by $f(q_U) - f_1(q_L + m_2^*)$, is positive when b is large since $q_U > q_L + m_2^*$. Since the felicity function is concave, increasing b increases the difference between the felicities. Thus this difference can be made arbitrarily large, indicating that $\widetilde{U}_1 > U_1^{provocation}$. The same argument can be made for provocation in which player 1 provokes by moving the policy to the upper bound. Thus for large enough b , provocation cannot be an equilibrium.

Proof of Lemma 1

First we observe that $\hat{q}_P(1) = \hat{q}_G(1)$. When there is only one firm, there is only one default policy in which the firm is indifferent between placation and pull. If the default policy is just slightly more to the left, placation is preferred, and if slightly more to the right, pull is preferred. We wish to show that when N is large, the equilibrium regions overlap for more than just a single point.

In order to show this, we must show that $\hat{q}_P(1) < \hat{q}_P(N)$ and $\hat{q}_G(N) < \hat{q}_G(1)$ for all $N > N'$. To prove the first inequality, we specify the condition in which no firm has an incentive to deviate from a placation equilibrium. This condition is given by⁶

$$f_F(q_U) - (\bar{q} - q_U) * c_F \geq f_F\left(\frac{N-1}{N}q_U + \frac{1}{N}\bar{q} + \tilde{m}_i(N) + m_G\left(\frac{N-1}{N}q_U + \frac{1}{N}\bar{q} + \tilde{m}_i(N)\right)\right) - N * \tilde{m}_i(N) * c_F \quad (\text{A.3})$$

where $\tilde{m}_i(N)$ is the optimal deviation. Notice that the optimal deviation will be moving the policy upward where marginal benefit equals marginal cost.

Let $\bar{q} = \hat{q}_P(1)$ so that expression (A.3) holds with equality when $N = 1$. When we allow $N > 1$, the left-hand side of equation (A.3) does not change. The right-hand side of the equation must decrease for two reasons. First, the cost function is now multiplied by N . Second, the other firms' actions place the default policy farther away from the bliss point and closer to q_U . Thus for (A.3) to continue to hold with equality, $\bar{q}$ must increase. This implies that $\hat{q}_P(N)$ is monotonically increasing in N .

Now we show that $\hat{q}_G(1) > \hat{q}_G(N)$ for all $N > N'$. The no deviation condition for a pull equilibrium is given by

$$f_F(\bar{q} + N * m_i^*(N) + m_G(\bar{q} + N * m_i^*(N))) - N * m_i^*(N) * c_F \geq f_F(q_U) - N * (\bar{q} - q_U + (N-1)m_i^*(N)) * c_F \quad (\text{A.4})$$

Notice that the best deviation is to placate. The equilibrium action already specifies the optimal action if the action moves the policy weakly higher. The only way that moving the policy lower could be optimal is if it placates.

When $\bar{q} = \hat{q}_G(N)$, then (A.4) holds with equality. Let $\bar{q} = \hat{q}_G(1)$. When N is very large, we will show that the right-hand side exceeds the left-hand side.

First, we show that $m_i^*(N) = 0$ for large N . The optimal action in pull, $m_i^*(N)$ will satisfy Player i 's first-order condition:

$$\frac{1}{N} f_F'(\bar{q} + N * m_i^*(N) + m_G(\bar{q} + N * m_i^*(N))) \left(1 + \frac{\partial m_G}{\partial m_i}\right) = c_F$$

Since the marginal benefit is decreasing by $\frac{1}{N}$ but the marginal cost is constant, there will be an N sufficiently large such that the marginal cost always exceeds the marginal benefit. This implies a corner solution to the first-order condition, hence $m_i^*(N) = 0$.

This means as N gets large, the left-hand side of (A.4) converges to $f_F(\bar{q} + m_G(0))$. The right-hand side explodes to negative infinity. Thus in order for this equation to hold with equality, it must be the case that $\bar{q}$ decreases, lowering the equilibrium utility on the left-hand side and increasing the deviation utility on the right-hand side. This implies that $\hat{q}_G(N) < \hat{q}_G(1)$ when N is large.

Proof of Proposition 2

Let us more carefully inspect the placation and pull equilibrium regimes as N increases so that we can determine how the equilibrium utility of these two equilibria changes. First let us look at equilibrium utility of a single firm in a symmetric placation equilibrium,

$$U_i(m_i, m_{-i}) = \frac{1}{N} f_F(q_U) - \frac{1}{N} (\bar{q} - q_U) * c_F$$

If we sum over all N firms we find that the total industry utility is invariant in N . This means that $U_{plac}(N; \bar{q})$ is invariant in N .

Remember from Lemma 1 that when $N > N'$ then $m_i^*(N) = 0$. This implies for $N > N'$, $U_{grav}(N; \bar{q}) = f_F(\bar{q} + m_G(\bar{q}))$. But clearly this expression is less than $U_{grav}(1; \bar{q}) = f_F(\bar{q} + m_1^*(1) + m_G(\bar{q} + m_1^*(1))) - m_1^*(1) * c_F$ since $m_1^*(1) = 0$ does not satisfy the firm's first-order condition. Thus from when $N = 1$ to when N is large, the industry equilibrium utility decreases.

Proof of Lemma 2

First we observe that $\hat{q}_R(1) = \hat{q}_{sl}(1)$. When there is only one firm, there is only one default policy in which the firm is indifferent between provocation and improvement. If the default policy is just slightly more to the left, provocation is preferred, and if slightly more to the right, improvement is preferred. We wish to show that when $N > 1$ the equilibrium regions overlap for more than just a single point.

In order to show this, we must show that $\hat{q}_R(1) < \hat{q}_R(N)$ and $\hat{q}_{sl}(1) > \hat{q}_{sl}(N)$ for all $N > 1$. To prove the first inequality, we specify the condition in which no firm has an incentive to deviate from provocation by extremists' equilibrium (the logic is the same as in the placation equilibrium):⁷

$$f_F(q_L - m_G(q_L)) - (\bar{q} - q_L) * c_F \geq f_F\left(\frac{N-1}{N} q_L + \frac{1}{N} \bar{q} + \tilde{m}_i(N)\right) - N * \tilde{m}_i(N) * c_F \quad (\text{A.5})$$

where $\hat{m}_i(N)$ is the optimal deviation. This deviation will move the policy upward toward the faction's bliss such that marginal costs equal marginal benefit.

Let $\bar{q} = \hat{q}_R(1)$ so that (A.5) holds with equality when $N = 1$. The left-hand side is invariant to N . The right-hand side of the equation, must decrease for two reasons. First, the cost function is now multiplied by N . Second, the other firms' actions place the default policy farther away from the bliss point and closer to q_L . Thus for (A.5), it must be the case that $\bar{q}$ increases, which decreases the equilibrium utility on the left-hand side and increases the deviation utility on the right-hand side. This implies that $\hat{q}_R(N)$ is monotonically increasing in N .

The logic we use here to show that $\hat{q}_{sl}(N) < \hat{q}_{sl}(1)$ is similar to the logic we used when examining the pull equilibrium. The condition for there to be no deviation from an improvement equilibrium is given by

$$\begin{aligned} f_F(\bar{q} + N * m_i^*(N)) - N * m_i^*(N) * c_F \geq \\ f_F(q_L + m_G(q_L)) - N * (\bar{q} - q_L + (N - 1)m_i^*(N)) * c_F \end{aligned} \quad (\text{A.6})$$

where $m_i^*(N)$ is the firm's equilibrium action in an improvement equilibrium. Notice that the best deviation is to induce a provocation equilibrium. The equilibrium action already specifies the optimal action if the action moves the policy weakly higher. The only way that moving the policy lower can be optimal is if it induces provocation.

When $\bar{q} = \hat{q}_{sl}(N)$ then expression (A.6) holds with equality. When $N > N'$ we will show the right-hand side exceeds the left-hand side. But first, we show the total optimal action $N * m_i^*(N)$ decreases when N is small. The optimal action in improvement, $m_i^*(N)$ will satisfy Player i 's first-order condition:

$$\frac{1}{N} f_{F'}(\bar{q} + N * m_i^*(N)) = c_F \Rightarrow m_i^*(N) = \max \left\{ \frac{f_F^{-1}(N c_F) - \bar{q}}{N}, 0 \right\}$$

Thus $m_i^*(N)$ strictly decreases with N until $m_i^*(N) = 0$ for all $N > N'$.

This means as N gets large, the left-hand side of (A.6) converges to $f_F(\bar{q})$ and the right-hand side explodes to negative infinity. This implies that $\hat{q}_{sl}(N) < \hat{q}_{sl}(1)$ when N is large.

Proof of Proposition 3

The following analysis parallels the proof for Proposition 2. First let us look at equilibrium utility of a single faction in the symmetric provocation by extremists' equilibrium:

$$U_i(m_i, m_{-i}) = \frac{1}{N} f_F(q_L - m_G(q_L)) - \frac{1}{N} (q_L - \bar{q}) * c_F$$

If we sum over all N factions we find that $U_{rev}(N; \bar{q})$ is invariant to N .

Remember from Lemma 2 that when $N > N'$ then $m_i^*(N) = 0$. This implies for $N > N'$, $U_{soup}(N; \bar{q}) = f_F(\bar{q})$. But clearly this expression is less than $U_{soup}(1; \bar{q}) = f_F(\bar{q} + m_1^*(1)) - m_1^*(1) * c_F$ since $m_1^*(1) = 0$ does not generically satisfy the faction's first-order condition. Thus total faction utility decreases when N is large.